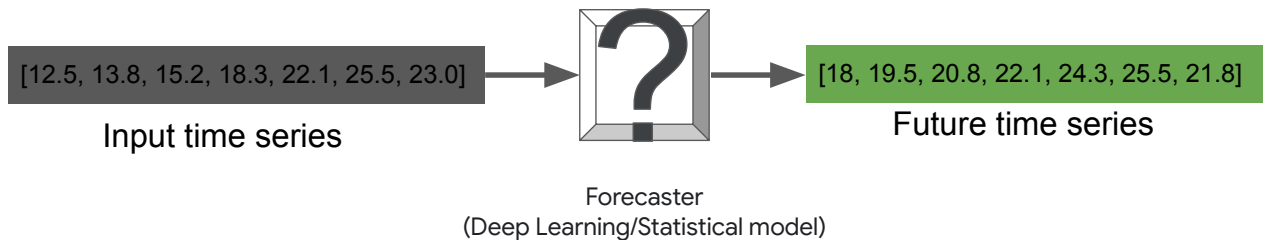


TFRBench: A Reasoning Benchmark for Evaluating Forecasting Systems

Md Atik Ahamed, Mihir Parmar, Palash Goyal, Yiwen Song, Long T. Le, Qiang Cheng, Chun-Liang Li, Hamid Palangi, Jinsung Yoon, Tomas Pfister

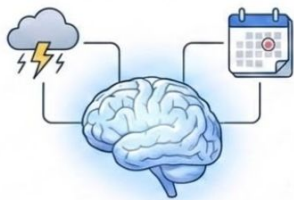
Time Series Forecasting in General



But what about context?
And what about reasoning?

The Gap in Reliability and Explainability

The “Why” Matters



A model might correctly predict a sudden spike in electricity demand, but it is much more valuable if it understands the cause, such as a specific weather event or a holiday.

The Challenge of Evaluation



Unlike numerical targets, there is no natural “ground truth” for the reasoning behind a forecast in existing benchmarks.

The Void



Currently, there are no standardized benchmarks that capture the reasoning quality behind a numerical forecast.

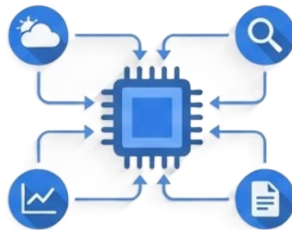
Introducing TFRBench

What is it?



The first benchmark specifically designed to evaluate the reasoning capabilities of forecasting systems alongside their numerical accuracy.

Defining “Reasoning”



In this context, reasoning is a model's ability to analyze multiple data channels, identify interdependencies, and provide helpful directive for future forecasting.

The Paradigm Shift



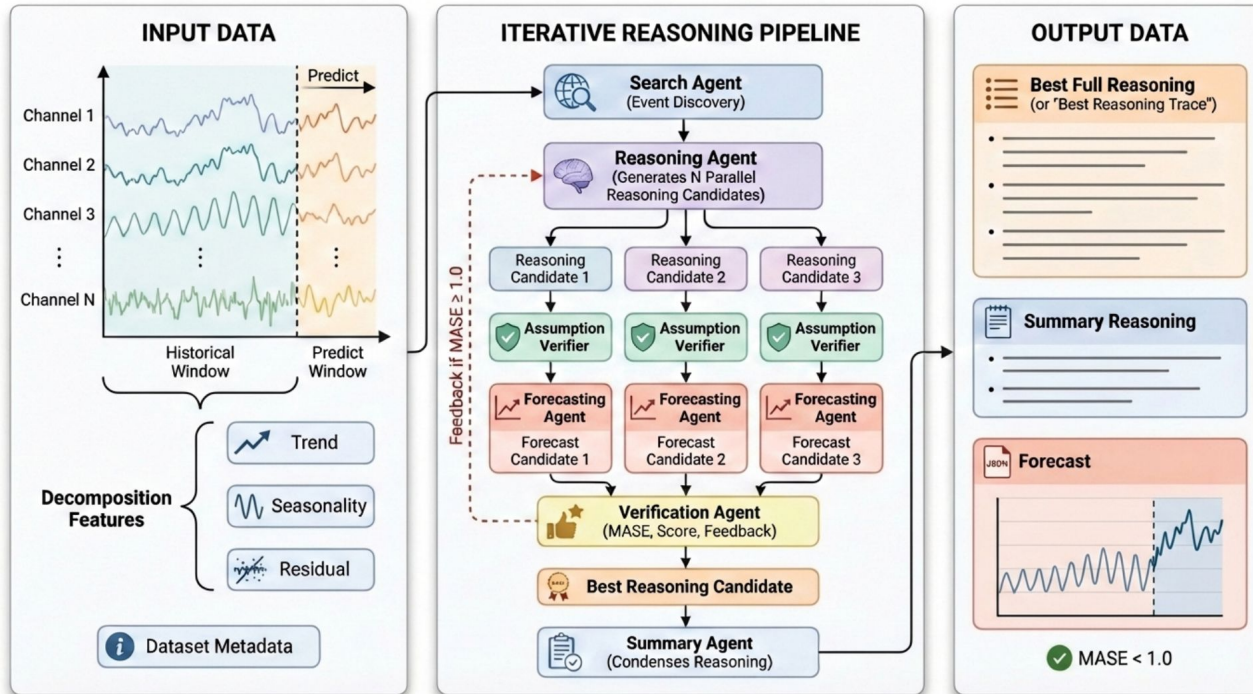
Moving from isolated numerical sequences to forecasts derived from explicit, verifiable causal logic.

The Scope of TFRBench

- **Broad Generalization:** Designed to test reasoning across varying data distributions.
- **Diverse Domains:** It spans 10 distinct datasets across 5 high-stakes domains.
- **Diverse temporal frequencies:** (Hourly to Monthly)
- **Univariate to Multivariate:** Expands from univariate (i.e., single channel) to multivariate (e.g., 5 channels)

Domain	Dataset	Source	Total Samples	Reference Samples	Freq.	Window (in-out)	Ch.	Series
Energy	Solar	GIFT-Eval	274	193	Daily	96-96	1	137
	Electricity	GIFT-Eval	1346	353	Daily	96-96	1	81
Sales	Car-parts	GIFT-Eval	1037	692	Monthly	25-25	1	1037
	Hierarchical Sales	GIFT-Eval	1370	1076	Daily	96-96	1	81
Web/CloudOps	Bitbrains - Fast Storage	GIFT-Eval	1153	678	Hourly	96-96	2	197
	Web Traffic	Monash archive	1214	710	Daily	96-96	1	162
Transportation	Traffic	Autoformer et al.	1104	449	Hourly	96-96	1	6
	NYC Taxi	nyc.gov	1000	391	Hourly	96-96	1	1
Economics/Finance	Amazon pricing	Yahoo Finance/nasdaq	512	328	Daily	14-7	5	1
	Apple pricing	Yahoo Finance/nasdaq	809	509	Daily	14-7	5	1

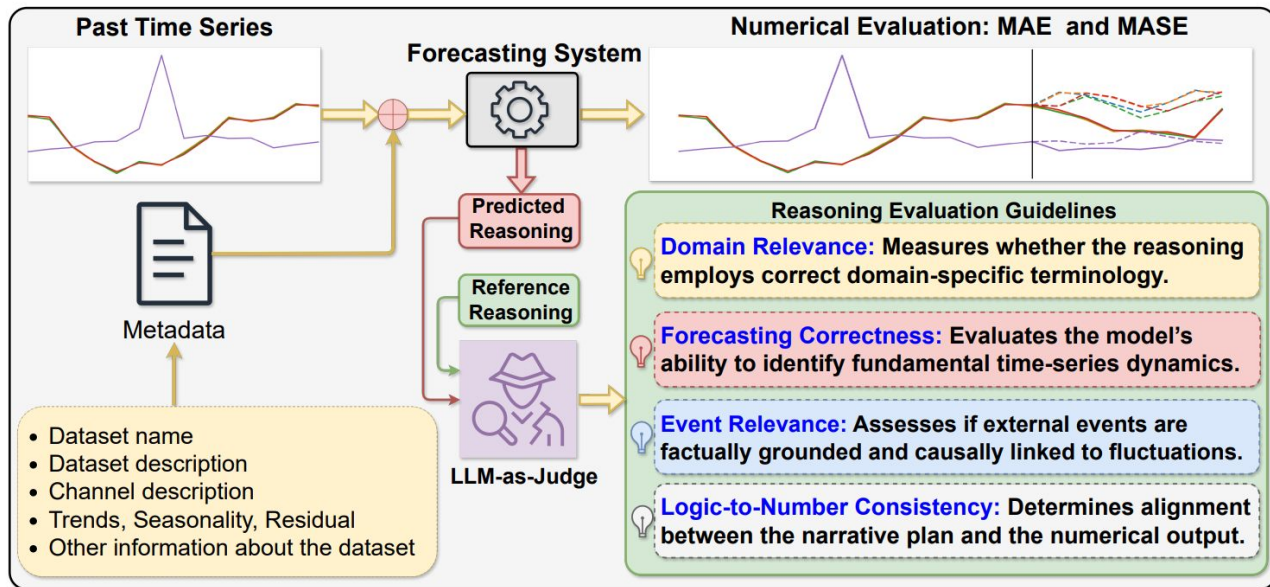
Overview of Data Creation System



Iterative "generate, verify, and refine" cycle

5 specialized agents: Search, Reasoning, Forecasting, Verification, Summary

Overview of the TFRBench Evaluation Pipeline



We employ an LLM-as-a-Judge to audit a candidate's reasoning trace against our Reference Reasoning.

The core focus of our framework is evaluating the quality of the reasoning.

Reasoning Evaluator (LLM-as-a-Judge)

Objective: Our goal is to evaluate the actual reasoning process behind forecasting. We use an LLM-as-a-Judge to audit candidates against “Reference Reasoning”, verifying their logical, factual, and domain-specific accuracy.

Domain Relevance	Forecasting Correctness	Event Relevance	Logic-to-Number Consistency
1. (Irrelevant/Wrong): Wrong domain terminology. Logic makes no sense. 2. (Generic): Vague language without domain terms. 3. (Acceptable): Basic terms used correctly but lacks depth. 4. (Good): Specific terminology and standard forecasting logic used correctly. 5. (Expert): Deep domain expertise, precise jargon, matches Ground Truth logic.	1. (Incorrect): Completely misses trend or seasonality. 2. (Weak): Identifies trend but ignores obvious seasonality. 3. (Average): Captures main trend/seasonality but misses magnitude/shape. 4. (Strong): Correctly identifies trend, seasonality, and general shape. 5. (Exact): Perfectly captures trend, seasonality, and inflection points.	1. (Hallucination): Mentions non-existent or false events. 2. (Irrelevant): Real events but no logical connection. 3. (Correlated): Identifies event but explanation is vague. 4. (Causal): Clearly links event to specific data movement. 5. (Aligned): Detailed, accurate causal explanation of impact.	1. (Contradiction): Text and numbers are opposites. 2. (Disconnected): Text describes shapes not seen in numbers. 3. (Weak Consistency): Direction matches but magnitude is off. 4. (Consistent): Numbers generally follow the narrative. 5. (Alignment): Numerical forecast is a precise translation of the reasoning.

LLM-as-a-Judge Results

Dataset	Models	w/ Reasoning				Event Forecast + Reasoning			
		Dom.	Fcst.	Evt.	Logic	Dom.	Fcst.	Evt.	Logic
Solar Daily	Gemini-2.5-Pro	3.026 _{0.000}	2.316 _{0.000}	2.036 _{0.000}	4.528 _{0.000}	3.803 _{0.000}	2.731 _{0.000}	2.207 _{0.000}	4.207 _{0.000}
	Gemini-3-Pro	4.699 _{0.000}	3.212 _{0.000}	2.508 _{0.000}	4.881 _{0.000}	4.119 _{0.000}	2.699 _{0.000}	2.368 _{0.000}	4.523 _{0.000}
	Claude-Sonnet-4.5	4.922 _{0.000}	3.062 _{0.000}	2.891 _{0.000}	3.363 _{0.000}	4.554 _{0.000}	2.637 _{0.000}	2.472 _{0.000}	3.368 _{0.000}
Electricity	Gemini-2.5-Pro	2.318 _{0.013}	2.768 _{0.020}	1.720 _{0.022}	4.590 _{0.015}	4.688 _{0.000}	3.076 _{0.000}	3.442 _{0.000}	3.198 _{0.000}
	Gemini-3-Pro	4.167 _{0.000}	3.173 _{0.000}	2.178 _{0.000}	4.773 _{0.000}	4.892 _{0.000}	3.606 _{0.000}	4.147 _{0.000}	4.484 _{0.000}
	Claude-Sonnet-4.5	2.054 _{0.000}	2.093 _{0.000}	2.167 _{0.000}	3.167 _{0.000}	4.609 _{0.000}	1.971 _{0.001}	2.885 _{0.001}	2.203 _{0.001}
Car Parts	Gemini-2.5-Pro	3.123 _{0.000}	2.658 _{0.000}	1.863 _{0.000}	4.181 _{0.000}	4.926 _{0.000}	3.464 _{0.000}	3.811 _{0.000}	4.923 _{0.000}
	Gemini-3-Pro	4.068 _{0.000}	2.900 _{0.000}	2.088 _{0.000}	4.945 _{0.000}	4.962 _{0.000}	3.012 _{0.000}	4.158 _{0.001}	4.883 _{0.000}
	Claude-Sonnet-4.5	4.120 _{0.000}	2.675 _{0.000}	2.022 _{0.000}	3.384 _{0.000}	4.939 _{0.000}	3.402 _{0.001}	4.084 _{0.000}	4.402 _{0.000}
Hierarchical Sales	Gemini-2.5-Pro	2.637 _{0.014}	2.972 _{0.007}	1.635 _{0.010}	4.730 _{0.015}	2.519 _{0.000}	2.456 _{0.000}	2.230 _{0.000}	3.346 _{0.000}
	Gemini-3-Pro	3.701 _{0.000}	2.878 _{0.000}	1.967 _{0.000}	4.335 _{0.000}	2.476 _{0.000}	2.492 _{0.000}	2.205 _{0.000}	3.798 _{0.000}
	Claude-Sonnet-4.5	3.757 _{0.000}	2.613 _{0.000}	2.046 _{0.000}	3.651 _{0.000}	2.198 _{0.000}	2.019 _{0.000}	1.947 _{0.000}	2.836 _{0.000}
Bitbrains Fast Storage	Gemini-2.5-Pro	3.807 _{0.000}	2.565 _{0.000}	2.944 _{0.000}	3.845 _{0.000}	3.802 _{0.000}	1.313 _{0.003}	2.652 _{0.001}	1.324 _{0.000}
	Gemini-3-Pro	4.531 _{0.000}	3.451 _{0.000}	3.249 _{0.000}	4.879 _{0.000}	4.746 _{0.000}	2.963 _{0.001}	2.704 _{0.006}	4.174 _{0.002}
	Claude-Sonnet-4.5	4.714 _{0.000}	2.780 _{0.000}	3.427 _{0.001}	2.507 _{0.000}	3.360 _{0.000}	1.432 _{0.000}	2.083 _{0.000}	1.543 _{0.000}
Web Traffic	Gemini-2.5-Pro	3.585 _{0.000}	2.831 _{0.000}	2.039 _{0.000}	4.693 _{0.000}	4.576 _{0.000}	2.761 _{0.000}	3.249 _{0.000}	3.645 _{0.000}
	Gemini-3-Pro	4.034 _{0.000}	3.169 _{0.000}	2.179 _{0.000}	4.904 _{0.000}	4.634 _{0.000}	2.804 _{0.000}	3.435 _{0.000}	4.699 _{0.000}
	Claude-Sonnet-4.5	3.928 _{0.000}	2.325 _{0.000}	2.120 _{0.000}	3.727 _{0.000}	4.276 _{0.000}	2.097 _{0.000}	2.669 _{0.000}	2.637 _{0.001}
Traffic	Gemini-2.5-Pro	2.762 _{0.012}	1.946 _{0.004}	1.896 _{0.027}	4.584 _{0.034}	3.643 _{0.001}	3.143 _{0.000}	2.203 _{0.000}	4.777 _{0.000}
	Gemini-3-Pro	3.793 _{0.000}	2.483 _{0.000}	2.249 _{0.000}	4.842 _{0.000}	2.604 _{0.000}	2.664 _{0.000}	2.022 _{0.000}	3.935 _{0.000}
	Claude-Sonnet-4.5	4.294 _{0.000}	2.254 _{0.000}	2.488 _{0.000}	3.586 _{0.000}	3.381 _{0.000}	2.379 _{0.000}	1.962 _{0.000}	3.454 _{0.000}
Nyc Taxi	Gemini-2.5-Pro	2.814 _{0.018}	1.922 _{0.010}	1.703 _{0.018}	4.360 _{0.040}	4.972 _{0.000}	3.529 _{0.001}	3.425 _{0.000}	4.652 _{0.000}
	Gemini-3-Pro	3.969 _{0.000}	2.113 _{0.000}	2.113 _{0.000}	4.872 _{0.000}	4.941 _{0.000}	3.844 _{0.000}	3.941 _{0.000}	4.555 _{0.000}
	Claude-Sonnet-4.5	4.532 _{0.000}	2.297 _{0.000}	2.363 _{0.000}	2.767 _{0.000}	4.890 _{0.000}	2.719 _{0.000}	3.156 _{0.000}	3.097 _{0.000}
Amazon	Gemini-2.5-Pro	2.262 _{0.000}	1.488 _{0.000}	1.637 _{0.000}	3.213 _{0.000}	4.966 _{0.000}	2.596 _{0.001}	3.617 _{0.001}	4.640 _{0.002}
	Gemini-3-Pro	2.936 _{0.000}	1.829 _{0.000}	1.784 _{0.000}	3.976 _{0.000}	4.997 _{0.000}	3.433 _{0.000}	4.125 _{0.000}	4.857 _{0.000}
	Claude-Sonnet-4.5	3.588 _{0.000}	2.012 _{0.000}	1.893 _{0.000}	3.970 _{0.000}	4.988 _{0.000}	2.543 _{0.000}	3.179 _{0.001}	4.491 _{0.000}
Apple	Gemini-2.5-Pro	2.217 _{0.001}	1.557 _{0.001}	1.695 _{0.000}	3.380 _{0.002}	4.986 _{0.000}	2.556 _{0.000}	3.648 _{0.000}	4.642 _{0.000}
	Gemini-3-Pro	3.000 _{0.000}	1.842 _{0.001}	1.800 _{0.000}	3.821 _{0.001}	5.000 _{0.001}	3.327 _{0.001}	4.313 _{0.001}	4.826 _{0.002}
	Claude-Sonnet-4.5	3.617 _{0.000}	2.075 _{0.000}	1.917 _{0.000}	3.786 _{0.000}	4.996 _{0.000}	2.515 _{0.000}	3.387 _{0.000}	4.642 _{0.001}

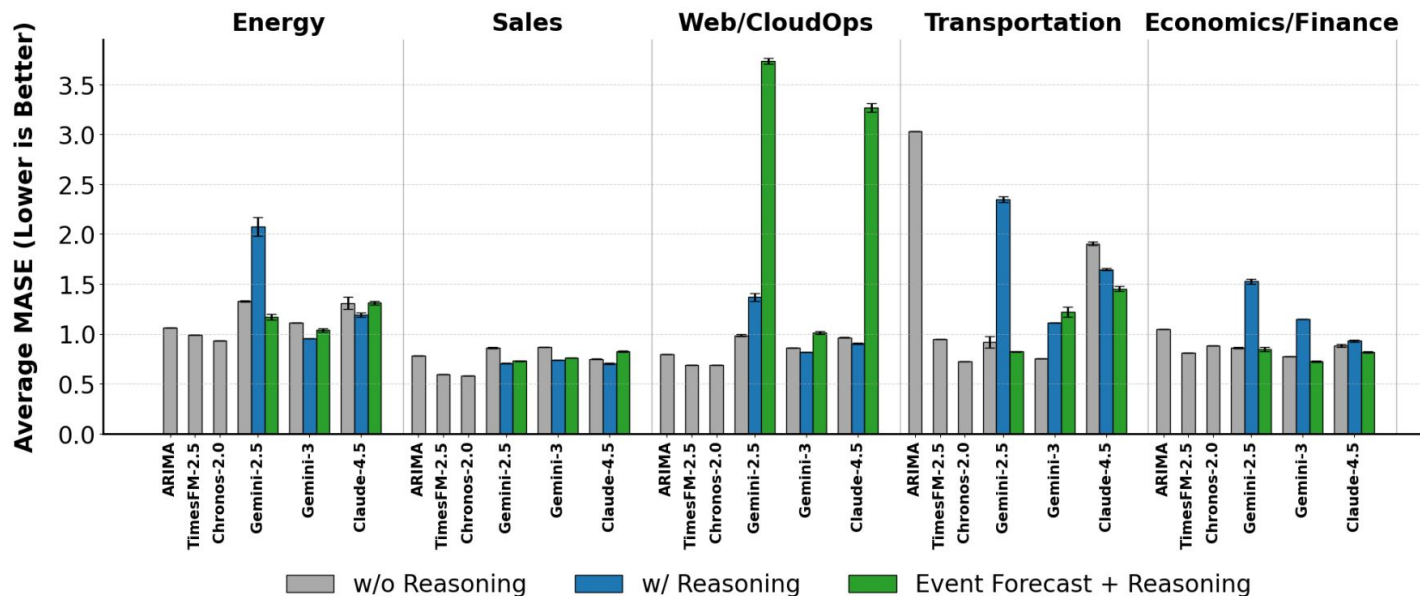
Majority of the models are not good at event relevance.

Though, they are good at logic to number conversion, their forecasting correctness are poor.

As forecasting correctness are poor, their reasoning is suffered from narrative bias.

This shows the risk of relying on regular LLM capabilities without the reasoning verification, TFRBench makes an effort to enable this.

Avg. MASE performance by domains

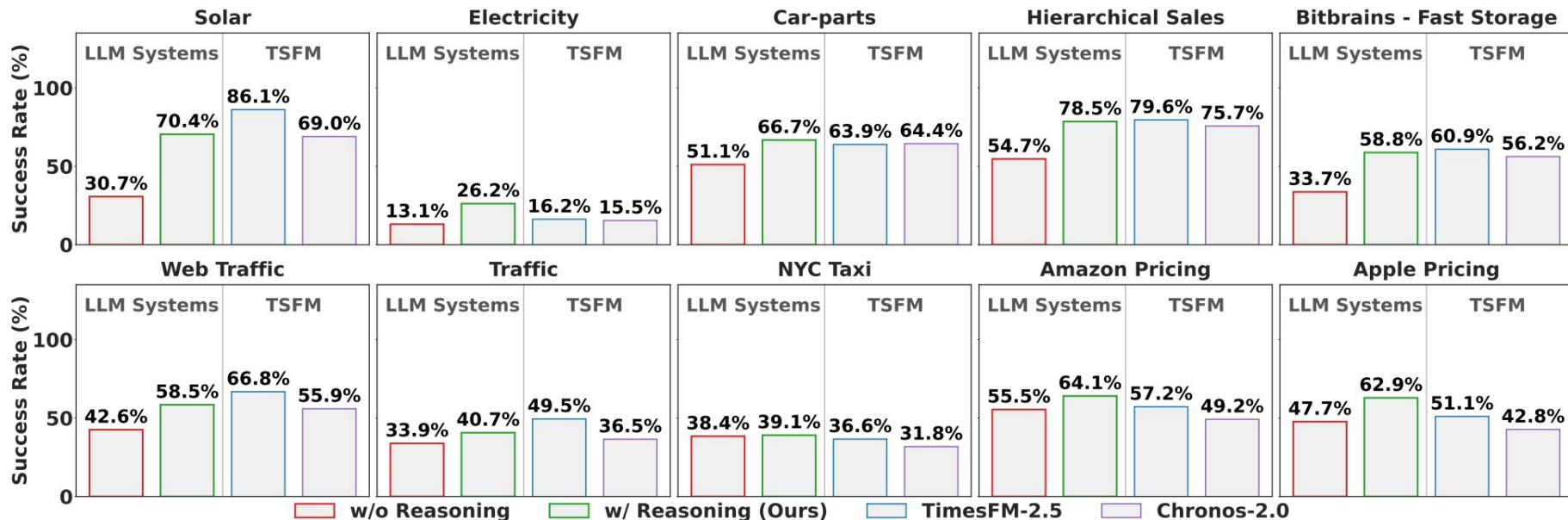


Accessing context for past data (e.g., weather, holidays) allows models to reason better, which is highly effective in Energy and Transportation.

However, this strategy is not always beneficial. In Web/CloudOps, event-aware reasoning acts as a distractor.

Key Finding: Reasoning Drives Accuracy

Overall Success Rate



➡ Prompting LLMs with our generated reasoning traces significantly improves forecasting accuracy compared to direct numerical prediction.

➡ LLMs are zero shot, while TimesFM-2.5 and Chronos-2.0 are pretrained with large numeric forecasting data.

Conclusion and Impact

A New Standard



TFRBench establishes a new benchmark for **interpretable, reasoning-based** evaluation in time-series forecasting.

Accountability



By requiring models to explain their predictions in high-stakes domains (Energy, Finance), we enhance human trust and accountability in automated decision-making.

Diagnostic Power



TFRBench serves as a vital diagnostic tool to identify biases and ensure forecasts are driven by factual logic, not statistical artifacts.